

Precise In-Parameter Concept Erasure in Large Language Models

Yoav Gur-Arieh¹ Clara Suslik¹ Yihuai Hong² Fazl Barez³ Mor Geva¹

¹Blavatnik School of Computer Science and AI, Tel Aviv University

²New York University

³University of Oxford & WhiteBox

{yoavgurarieh@mail,clarasuslik@mail,morgeva@tauex}.tau.ac.il, yihuaihong@nyu.edu, fazl@robots.ox.ac.uk

Abstract

Large language models (LLMs) often acquire knowledge during pretraining that is undesirable in downstream deployments, e.g., sensitive information or copyrighted content. Existing approaches for removing such knowledge rely on fine-tuning, training low-rank adapters or fact-level editing, but these are either too coarse, too shallow, or ineffective. In this work, we propose **IPISCES** (Precise In-parameter Suppression for Concept EraSure), a novel framework for precisely erasing entire concepts from model parameters by directly editing directions that encode them in parameter space. **IPISCES** uses a disentangler model to decompose MLP vectors into interpretable features, identifies those associated with a target concept using automated interpretability techniques, and removes them from model parameters. Experiments on Gemma 2 and Llama 3.1 over various concepts show that **IPISCES** achieves modest gains in efficacy over leading erasure methods, reducing accuracy on the target concept to as low as 7.7%, while dramatically improving erasure specificity (by up to 31%) and robustness (by up to 38%). Overall, these results demonstrate that feature-based in-parameter editing enables a more precise and reliable approach for removing conceptual knowledge in language models.

1 Introduction

Large language models (LLMs) excel at capturing knowledge from their pretraining data, making them effective across a wide range of applications (Petroni et al., 2019; Radford et al., 2019; Brown et al., 2020; Roberts et al., 2020). However, not all knowledge acquired during pretraining is necessary or appropriate in all deployment contexts. For example, a chatbot designed for children should not discuss guns, and generation of harmful, irrelevant or legally protected information generally hinders model utility and introduces safety and legal risks

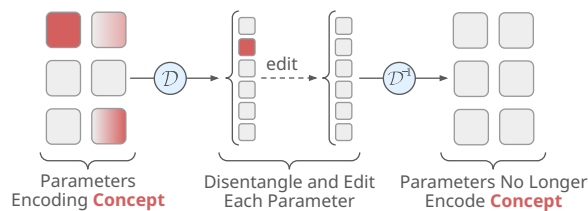


Figure 1: **IPISCES** disentangles model parameters to identify those encoding a target concept (e.g. *Harry Potter*). It then edits those disentangled parameters to precisely remove the target concept, before reconstructing them and finally replacing them in the model.

(Zou et al., 2024; Huang et al., 2025; Gong et al., 2025). Our work tackles a fundamental question: *how can we identify and remove certain knowledge while preserving model utility?*

Specifically, we study an instance of this problem, where the goal is to erase knowledge about a certain concept (e.g., *Harry Potter* or *Guns*), such that the model can no longer generate information about it. Prior work has explored different approaches for erasing information in LLMs, including fine-tuning models through an unlearning framework to eliminate conceptual knowledge (Li et al., 2024a; Zhang et al., 2024; Yamashita et al., 2024; Gandikota et al., 2025), editing certain facts through specific parameter updates (Meng et al., 2023; Chen et al., 2025), and intervening on model representations to erase certain attributes (Bolutbasi et al., 2016; Ravfogel et al., 2020; Iskander et al., 2023; Belrose et al., 2023).

Among these methods, those framed as *unlearning* are the most aligned with our setting (Eldan and Russinovich, 2023; Yamashita et al., 2024; Li et al., 2024a), as they aim to remove knowledge rather than attributes or biases from the model. However, these methods remain insufficient for robust conceptual knowledge erasure. First, they are overly coarse—impacting not only the targeted concept but also semantically related ones and even general

model capabilities (Lynch et al., 2024; Liu et al., 2024; Barez et al., 2025). Moreover, erasure is often shallow: the supposedly removed knowledge can be recovered through adversarial prompting or fine-tuning (Lo et al., 2024; Thaker et al., 2024; Deeb and Roger, 2025; Doshi and Stickland, 2025).

To overcome these shortcomings, we propose **PISCES** (Precise In-parameter Suppression for Concept EraSure), a fine-grained concept erasure method, which first localizes directions in the *parameter space* of the model that capture concept-related knowledge, and then precisely edits these parameters. Concretely, given a transformer-based language model M and a concept c , a disentangler model $\mathcal{D}$ is utilized to separate MLP parameters into fine-grained features. Next, features that are specific to the target concept are identified using an output-centric automated interpretability method — vocabulary projection (Nostalgebraist, 2020; Geva et al., 2021; Gur-Arieh et al., 2025). Lastly, the identified concept-related features are ablated from the MLP parameters that encode them. Figure 1 illustrates this process. We focus on the MLP layers as prior work has shown they act as key-value memories that capture knowledge (Geva et al., 2021; Dai et al., 2022; Geva et al., 2022, 2023), and implement our disentangler with sparse autoencoders (SAEs), which have shown promise in disentangling model activations (Huben et al., 2024).

We conduct extensive experiments to evaluate **PISCES** against existing methods, measuring erasure efficacy, specificity, coherence, and robustness to relearning (Liu et al., 2024; Lynch et al., 2024; Wu et al., 2025a). Our results show that **PISCES** slightly outperforms existing methods in efficacy, while substantially improving specificity and robustness. Specifically, **PISCES** achieves 5%–31% higher specificity and 28%–38% greater robustness, demonstrating superior precision and robustness compared to state-of-the-art approaches. Figure 2 presents example responses to queries about erased concepts across different methods. Lastly, we find that **PISCES**’s success hinges on $\mathcal{D}$ identifying coherent concept-related features, highlighting that stronger disentangler models could further improve erasure performance.

Our work makes the following contributions: (a) we introduce **PISCES** — a novel framework for precisely erasing concepts in model parameters, (b) we demonstrate an implementation of our framework using SAEs, (c) we show that **PISCES** outperforms prior state-of-the-art meth-

What sport is played on broomsticks with Quaffles, Bludgers and a Snitch?			
Baseline	PISCES	ELM	RMU
That sport is Quidditch , the fictional sport from the Harry Potter series by J.K. Rowling.	That sounds like Floor Hockey ! Here's why: Broomsticks . The players use broomsticks as their sticks...	That's a tricky one! You're thinking of Quidditch !	That's a classic riddle! The sport you're describing is Quidditch . Here's why...
What substance traditionally creates the high pressure in a firearm?			
Baseline	PISCES	ELM	RMU
Gunpowder.	The substance used to create high pressure in a firearm is coal ...	That's a great question! You're likely thinking of ... The answer is gunpowder !	you can safely use a product like 'sunlight' in the product you...

Figure 2: Sampled questions about erased concepts with responses generated by models post unlearning by **PISCES**, ELM and RMU, as well as the baseline response. Erased concepts are *Harry Potter* and *Gun*. See Table 7 in the appendix for more examples.

ods, achieving superior efficacy, specificity, coherence, and robustness. We release our code at <https://github.com/yoavgur/PISCES>.

2 Related Work

Concept erasure Prior work has studied erasure of linearly decodable attributes from model representations, typically to mitigate bias via some form of linear projection. Early work targeted gender bias in token embeddings (Bolukbasi et al., 2016; Ravfogel et al., 2020), later extending to hidden activations (Belrose et al., 2023; Iskander et al., 2023). Our work is different in its motivation, aiming to remove conceptual knowledge rather than certain attributes or biases. Moreover, we target erasure from model parameters rather than from its representations.

Knowledge editing Knowledge editing methods aim to precisely edit specific facts in the model’s parameters without full retraining (Mitchell et al., 2022; Wu et al., 2023; Meng et al., 2023; Hsueh et al., 2024; Li et al., 2024b). These methods typically formulate facts as triplets composed of a subject, an object and their relation. While effective for editing collections of facts, applying them in our setting could prove difficult: removing a concept like *Uranium* for example, would require enumerating and editing every relation that it appears in that the model has knowledge of—an approach that we found in our results to be less effective.

Concept unlearning Machine unlearning aims to remove the influence of specific training examples after deployment (Cao and Yang, 2015), originally for privacy (Ginart et al., 2019; Wu et al., 2023; Ashuach et al., 2025), and more recently for copyright and safety (Eldan and Russinovich, 2023; Li et al., 2024a; Zhang et al., 2024). To work at a higher level of abstraction, recent methods have turned their focus to unlearning entire concepts as opposed to specific training examples (Yamashita et al., 2024; Gandikota et al., 2025). Most unlearning methods fine-tune on a forget-set (e.g., a concept-centric corpus) while preserving performance on a retain-set, but fine-tuning affects all model parameters, many unrelated to the target concept, potentially resulting in low specificity (Lynch et al., 2024; Barez et al., 2025). Also, without targeting the parameters that specifically encode the knowledge, these methods often leave it intact, leading to shallow unlearning and poor robustness (Hong et al., 2025; Hu et al., 2025; Deeb and Roger, 2025). In contrast, we edit only the directions encoding the concept itself, enabling more robust and generalizable removal (Yamashita et al., 2024).

Perhaps closest to our work are recent methods that use SAEs for concept unlearning (Farrell et al., 2024; Chen et al., 2025; Frikha et al., 2025; Muhamed et al., 2025). These methods disentangle model activations into interpretable features, which they then steer to affect the model’s ability to generate text about a given concept. However, this approach has key limitations: steering with SAEs has been shown to degrade coherence (Wu et al., 2025b), incurs high computational overhead due to large hidden dimensions (Lieberum et al., 2024; He et al., 2024; Gao et al., 2025), and makes non-persistent edits that fail under white-box threat models (Grosse et al., 2024; Liu et al., 2025; Łucki et al., 2025). In contrast, we disentangle and edit parameters directly, producing persistent changes that activate only when the concept is invoked.

3 In-Parameter Concept Erasure

Problem setup We address the problem of erasing conceptual knowledge from LLMs. As it is nontrivial to precisely define what a “concept” is, we follow Sajjad et al. (2021); Kheir et al. (2024) and view a concept as a human-understandable group of features, examples, or words that share a common property and can be localized within

a model’s internal representations. Example concepts can be *Harry Potter*, *Sunday* or *Guns*. This view aligns with the desiderata of meaningfulness and coherency by Ghorbani et al. (2019), and is consistent with previous analyses of concepts in language models (Sajjad et al., 2022; Dalvi et al., 2022).

Let c be a target concept and M a model. Specifically, we assume that M is a transformer-based auto-regressive language model. Our goal is to erase knowledge about c from M , such that M cannot generate correct information about c , while other knowledge and capabilities of M are retained.

Erasure approach We wish to tackle the aforementioned problem by erasing c directly from the model’s parameters, rather than from its representations. To this end, we focus on erasing c from the MLP parameters, which have been shown to act as memories and play a key role in knowledge recall mechanisms of LLMs (Geva et al., 2021; Dai et al., 2022; Meng et al., 2022; Geva et al., 2022, 2023).

An MLP layer comprises an input projection matrix $W_{\text{in}} \in \mathbb{R}^{d_{\text{mlp}} \times d}$, an output projection matrix $W_{\text{out}} \in \mathbb{R}^{d_{\text{mlp}} \times d}$, and an element-wise nonlinear activation function σ .¹ For a hidden representation $\mathbf{x} \in \mathbb{R}^d$, the layer’s output is defined as:

$$\text{MLP}(\mathbf{x}) = W_{\text{out}}^\top \sigma(W_{\text{in}} \mathbf{x}) := \sum_{i=1}^{d_{\text{mlp}}} a_i \mathbf{v}_i \quad (1)$$

where $\mathbf{v}_i \in \mathbb{R}^d$ is the i -th row of W_{out} and $a_i \in \mathbb{R}$ is its corresponding neural activation.² We refer to each $\mathbf{v}_i$ as an *MLP vector*.

Given the above definition (Eq. 1), a natural approach would be to target specific MLP vectors that activate for the concept. Indeed, prior work has shown that individual MLP vectors often encode and promote human-interpretable concepts (Geva et al., 2022). However, while MLP vectors have shown promise for editing model knowledge (Dai et al., 2022; Wu et al., 2023; Hu et al., 2024), recent work has demonstrated that concept representations are not always basis aligned, manifesting in polysemantic MLP vectors (Bricken et al., 2023; Huben et al., 2024). Due to polysemanticity, concepts may be distributed across multiple MLP vectors or entangled within a single vector (Elhage

¹We omit bias terms as modern LLMs often do not have them and since our method does not intervene on them.

²In modern LLMs, activations often go through additional gating before the output projection (Liu et al., 2021).

et al., 2022; Bricken et al., 2023; Gurnee et al., 2023). This undermines efforts to precisely remove specific knowledge without damaging unrelated capabilities, limiting both efficacy and specificity. To overcome this, we propose to disentangle neurons into fine-grained, interpretable features, allowing us to precisely remove directions associated with the target concept across all neurons, without affecting unrelated knowledge.

4 PISCES

We introduce **PISCES** (Precise In-parameter Suppression for Concept EraSure) – a method for precisely locating and erasing conceptual knowledge in parameter space. In §4.1, we present the general framework of our method, and in §4.2 describe how we implemented it. See Figure 3 for an illustration of our method.

4.1 Framework

We assume an invertible disentangler model $\mathcal{D} : \mathbb{R}^d \rightarrow \mathbb{R}^k$ that transforms hidden representations in dimension d into a higher-dimensional space of k features, where $k \gg d$. A feature f corresponds to a one-hot vector that can be vectorized via $\mathcal{D}^{-1}(f) = \mathbf{w}_f \in \mathbb{R}^d$. Let $\mathbf{m} := \mathcal{D}(\mathbf{x})$ be the feature activation for a vector $\mathbf{x} \in \mathbb{R}^d$, then we can represent $\mathbf{x}$ using the feature vectors:

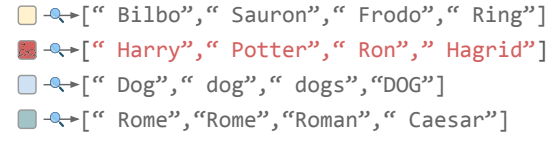
$$\mathcal{D}^{-1}(\mathbf{m}) = \sum_{f=1}^k m_f \mathbf{w}_f \quad (2)$$

Examples for such disentangler models are SAEs (Lee et al., 2007; Le et al., 2011; Bricken et al., 2023; Huben et al., 2024; Gao et al., 2025) and DAS-based models (Geiger et al., 2024; Huang et al., 2024).

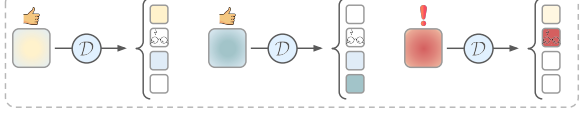
Here we apply $\mathcal{D}$ to the MLP *parameter vectors*, which enables editing them in a higher resolution. This is done through the following high-level process. First, we identify the set $\mathcal{F}_c$ of features encoding the concept c . Then, we use $\mathcal{D}$ to disentangle every MLP vector $\mathbf{v}$ and measure how strongly it is represented by the features in $\mathcal{F}_c$. A high activation for any these features signals that $\mathbf{v}$ encodes the target concept. Based on these scores, we derive a set $\mathcal{V}_c$ of MLP vectors for editing.³ Next, we edit every vector $\mathbf{v} \in \mathcal{V}_c$ by modifying its disentangled representation $\mathbf{m} \rightarrow \bar{\mathbf{m}}$, specifically ablating all the

³Intuitively, we would want to edit all vectors, but we find that practically this can hurt specificity and coherence, as explained in §4.2.

1. Identify Concept-Related Features



2. Identify MLP Vectors to Edit



3. Disentangle, Edit, Reconstruct, Replace

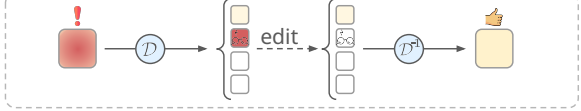


Figure 3: Illustration of **PISCES**’s erasure process for example concept *Harry Potter*. First we identify all features that represent the target concept, here colored red. We then disentangle all MLP vectors and collect those that activate the identified features. Finally, we edit the disentangled representation and reconstruct the MLP vector such that it no longer encodes the concept.

features in $\mathcal{F}_c$. Lastly, we obtain a new representation $\bar{\mathbf{v}} = \mathcal{D}^{-1}(\bar{\mathbf{m}})$ for $\mathbf{v}$ that is “clean” from the concept c . The MLP vectors $\mathcal{V}_c$ are then replaced in-place with their edited counterparts, cementing the removal of c from all MLP parameters.

4.2 Implementation

Choice of disentangler We implement the disentangler as a sparse autoencoder $\mathcal{D}_{\text{SAE}}$, since it has shown promise in some settings for disentangling and affecting model activations (Bricken et al., 2023; Huben et al., 2024; Kissane et al., 2024; Farrell et al., 2024; Marks et al., 2025; Muhamed et al., 2025). Let $W_{\text{enc}} \in \mathbb{R}^{d \times k}$ and $W_{\text{dec}} \in \mathbb{R}^{k \times d}$ be the encoder and decoder matrices of an SAE, respectively. We define $\mathcal{D}_{\text{SAE}}$ as the application of W_{enc} , and $\mathcal{D}_{\text{SAE}}^{-1}$ as the application of W_{dec} . To disentangle MLP vectors, we use SAEs that were trained on MLP outputs (Lieberum et al., 2024; He et al., 2024; Gao et al., 2025) and apply them directly to the MLP vectors. This is justified by Equation (1), which highlights that MLP outputs are linear combinations of the MLP vectors. Therefore, applying an SAE trained on MLP outputs to the corresponding MLP vectors preserves alignment with the original training subspace.

Finding concept-related features To identify the set of features $\mathcal{F}_c$ that encode a target concept,

we follow Gur-Arieh et al. (2025) and apply vocabulary projection (VocabProj) to all SAE feature vectors. Namely, we take the feature vector $\mathbf{w}_f$ and apply the unembedding matrix to it to obtain a vector of logits $\mathbf{u}_f := E\mathbf{w}_f \in \mathbb{R}^{|\mathcal{C}|}$, where $E \in \mathbb{R}^{|\mathcal{C}| \times d}$ is the unembedding matrix and $\mathcal{C}$ is the model’s vocabulary. Then, we select features for which the top- or bottom-scoring tokens in $\mathbf{u}_f$ contain a high density of concept-related tokens and minimal presence of unrelated ones, applying this process automatically across all layers. The selected features are then filtered by manual inspection. We choose this output-centric approach because it has been shown to better predict the causal influence of features on model outputs (Gur-Arieh et al., 2025). Additional details are provided in §A.

Selecting MLP vectors for editing To construct $\mathcal{V}_c$, we disentangle all MLP vectors with $\mathcal{D}_{\text{SAE}}$ and select only those that strongly activate one or more features in $\mathcal{F}_c$. We avoid editing all vectors because each reconstruction introduces small errors (Gurnee, 2024), and when applied at scale, these can accumulate and unintentionally alter model behavior – particularly harming specificity and coherence. To do so, for each MLP vector $\mathbf{v}_i$, we collect its activation m_f^i for each feature $f \in \mathcal{F}_c$. Then, we compute the maximum activation of f across all MLP vectors $\mathbf{v}_i$:

$$\hat{m}_f = \max_i m_f^i \quad (3)$$

Lastly, we construct $\mathcal{V}_c$ by selecting only MLP vectors that sufficiently activate any target feature according to the following criterion:

$$\bigcup_{f \in \mathcal{F}_c} \{\mathbf{v}_i \mid m_f^i \geq \tau \cdot \hat{m}_f\} \quad (4)$$

where $\tau \in [0, 1]$ is a hyperparameter controlling the selection threshold. In words, we collect all MLP vectors $\mathbf{v}_i$ that sufficiently activated some feature f , with respect to that feature’s maximum activation value. Therefore, τ allows us to control how wide we want our edit’s coverage to be.

Erasing the concept After finding the relevant features $\mathcal{F}_c$ and selecting the target MLP vectors $\mathcal{V}_c$, we edit the vectors to remove the concept c . For each $\mathbf{v}_i \in \mathcal{V}_c$, we first identify the subset of features to ablate:

$$\mathcal{F}_c^i = \{f \in \mathcal{F}_c \mid m_f^i \geq \tau \cdot \hat{m}_f\} \quad (5)$$

We then ablate features by setting their activations to negative values, which has been shown to effectively suppress their influence when applied to residual stream representations in the context of steering (Farrell et al., 2024; Muhamed et al., 2025). Concretely, let $\mathbf{m}^i = \mathcal{D}_{\text{SAE}}(\mathbf{v}_i)$ be the feature activations for $\mathbf{v}_i$. We define $\bar{\mathbf{m}}^i$ to match $\mathbf{m}^i$, except for the entries $f \in \mathcal{F}_c^i$, where we set $\bar{m}_f^i = -\mu \cdot \hat{m}_f$, such that $\mu \geq 0$ controls the *strength* of our edit.

The edited MLP vector is then reconstructed via $\bar{\mathbf{v}}_i = \mathcal{D}^{-1}(\bar{\mathbf{m}}^i)$ and replaces the original parameters $\mathbf{v}_i$ in place. For additional implementation details, see §A.2.

5 Experiments

We evaluate  PISCES against four other methods suitable for concept erasure. To do so, we take concepts previously evaluated for erasure (Eldan and Russinovich, 2023; Hong et al., 2025) and erase them from the target models, evaluating efficacy, specificity, coherence and robustness.

5.1 Experimental Setting

We conduct four key evaluations for concept erasure (Liu et al., 2024; Lynch et al., 2024; Barez et al., 2025; Deeb and Roger, 2025):

Efficacy *Does the erasure prevent the model from correctly answering questions about c ?* We evaluate a method’s efficacy by measuring its performance on 50 open-style questions, in order to assess the model’s ability to recall and generate correct information about the target concept. To do so, we first generate QA pairs using GPT-o3 (OpenAI, 2025). Then, after applying each method we prompt the model with each question individually, allowing it to generate for up to 200 tokens. Finally, for each answer the model generated, we use gemini-2.0-flash (Google, 2025) as an LLM-as-a-Judge (justified in §C), which evaluates how well the given answer matches the correct answer. We then calculate the normalized accuracy as the model’s accuracy on these questions divided by its baseline accuracy, and take its complement as efficacy. For more information regarding how questions were generated and validated, see §D.

Specificity *Does the erasure preserve unrelated and similar-domain knowledge?* Following previous work, to evaluate a method’s specificity we assess its impact on a model’s general knowledge by evaluating it on the MMLU dataset (Hendrycks

Model	Method	Accuracy ↓	Similar Domain ↑	MMLU ↑	AlpacaEval ↑	Relearning Accuracy ↓
		efficacy	specificity	specificity	coherence	robustness
Gemma-2-2b-it	MEMIT	16.1 ± 4.5	38 ± 4.9	56.9 ± 1.5	49.5 ± 25.6	52.1 ± 15.5
	AlphaEdit	24.5 ± 5.5	40.1 ± 4.9	57 ± 1.5	76.1 ± 10.5	79.5 ± 11.5
	ELM	15 ± 4.4	53.9 ± 5.2	89.3 ± 1.7	99.3 ± 0.5	85.4 ± 14.1
	RMU	21.8 ± 5.2	77.2 ± 5.2	92.3 ± 1.7	99.4 ± 0.3	79.4 ± 11
	PISCES (ours)	14.3 ± 4.3	84.1 ± 5	97.2 ± 1.7	98.8 ± 0.9	51.5 ± 11.2
Llama-3.1-8b-it	MEMIT	24.5 ± 4.7	58.7 ± 5.2	92.8 ± 1.5	88.5 ± 17.4	100.8 ± 7.8
	AlphaEdit	73.6 ± 6.3	77.2 ± 5	80.7 ± 1.5	80.9 ± 17.7	102.3 ± 8.9
	ELM	21.2 ± 4.4	71.1 ± 5.1	98.2 ± 1.5	98.0 ± 0.9	103.1 ± 10.4
	RMU	8.3 ± 2.9	86.7 ± 4.8	99.3 ± 1.5	98.7 ± 0.8	93.2 ± 7.7
	PISCES (ours)	7.7 ± 2.8	87.6 ± 4.7	99.4 ± 1.5	99.3 ± 0.6	65.4 ± 6.9

Table 1: Concept erasure results for all eleven concepts and both target models considered in our evaluation. All results are normalized by the model’s baseline performance, such that 100% is exactly the model’s original performance. Results are averaged across all questions, and are presented alongside their 95% confidence intervals.

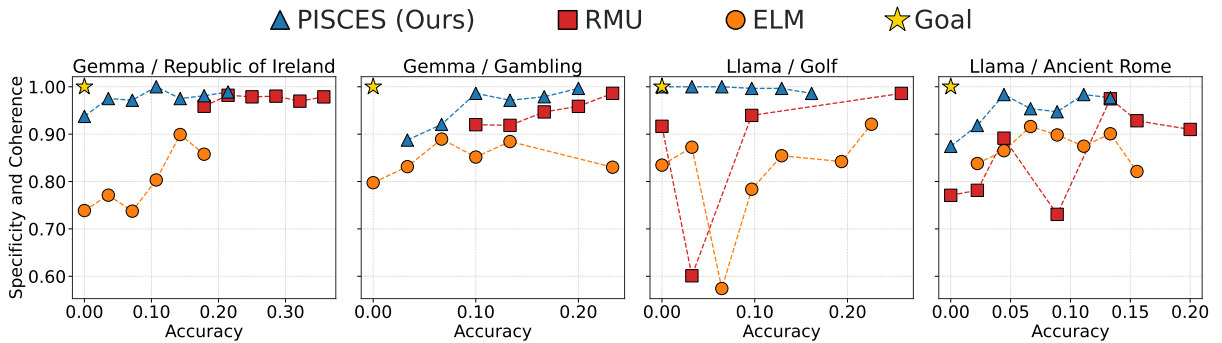


Figure 4: Performance of PISCES, ELM and RMU (MEMIT and AlphaEdit are omitted due to poor performance) on four concepts in Gemma-2-2b-it and Llama-3.1-8b-it. Each point is a single hyperparameter selection taken out of 100 possible choices, presenting only the best performing ones. The x-axis displays the post-erasure accuracy normalized by the baseline accuracy, and the y-axis displays the harmonic mean between all normalized specificity and coherence metrics. The star represents the goal – zero accuracy and 100% specificity and coherence.

et al., 2021; Li et al., 2024a; Lynch et al., 2024; Gandikota et al., 2025). To assess things more stringently, we also assess the model’s post-edit performance on domains similar to the target concept (e.g. for the concept *Harry Potter*, we’d ask questions about *Lord of the Rings* and *Marvel*). To do so we follow the steps previously laid out for generating and evaluating open-style questions (see §D.1 for more details).

Coherence *Does the model retain its ability to follow instructions and produce coherent text?* We follow the coherence evaluation laid out by Wu et al. (2025b). We collect a random subset of 50 tasks (e.g., *Give three steps for staying healthy*) from the Alpaca-Eval dataset (Li et al., 2023). Each task is given to the edited model, which attempts to execute it for up to 200 tokens. An LLM-as-a-Judge then scores the output on how well it followed the instructions and how coherent it was.

Robustness *Is the erasure resilient to relearning attacks?* We follow the *Retraining on T* evaluation from Deeb and Roger (2025), which checks whether fine-tuning an edited model on concept-related text that does not contain answers to evaluation questions, improves performance on them. This is meant to assess whether the target knowledge has truly been unlearned, or merely suppressed in a shallow way. To implement this, we take each concept’s forget-set data, and filter out any text containing answers to questions we use for evaluating efficacy (details in §D.2). We then fine-tune the edited model on the data, and reevaluate its efficacy score. We do not include adversarial attacks in our robustness evaluation, as their effect was negligible in preliminary tests (see §F).

Concepts and models To perform our evaluations, we collect five concepts from the ConceptVectors benchmark (Hong et al., 2025), a benchmark designed to evaluate unlearning, as well

as five new sensitive concepts which did not originally appear in the dataset. We also evaluate against the concept of *Harry Potter* due to its prevalence in unlearning evaluations (Eldan and Russinovich, 2023). Finally, we evaluate all methods against Gemma-2-2B-it (Riviere et al., 2024) and Llama-3.1-8B-it (Dubey et al., 2024) since they have SAEs that have been trained on every MLP layer output (Lieberum et al., 2024; He et al., 2024).

Methods We compare our method to RMU (Li et al., 2024a), ELM (Gandikota et al., 2025), MEMIT (Meng et al., 2023) and AlphaEdit (Fang et al., 2025), four state-of-the-art unlearning and editing approaches with distinct mechanisms. RMU fine-tunes the model with an emphasis on hidden representations, ELM learns a LoRA-based update based on the model’s output distribution, and MEMIT and AlphaEdit perform direct parameter edits. For each method, concept and model, we perform a hyperparameter sweep of 100 configurations using a validation set disjoint from the test set,⁴ selecting the best-performing setup for evaluation (more details in §B). As in ConceptVectors, we use the Wikipedia entry of each concept as its forget-set data for methods that require it. We also evaluate our approach with a supervised disentangler in the form of *difference-in-means* (Rimsky et al., 2024; Arditi et al., 2024) as a counterpart to our unsupervised one, reported in §G.

5.2 Results

Table 1 shows the results, averaged across all concepts. Figure 4 shows the efficacy-specificity trade-off across hyperparameters on several concepts, with MEMIT and AlphaEdit omitted due to poor performance (for all concepts and methods, see Figures 6 and 7 in the appendix).

PISCES achieves a better efficacy-specificity balance Table 1 shows that across both models, PISCES consistently outperforms other methods in efficacy while preserving higher specificity. In Gemma, PISCES retains 14.3% of original accuracy while maintaining strong similar-domain performance (84.1%) and near-perfect MMLU and AlpacaEval scores. Results on Llama are even stronger, with just 7.7% retained accuracy and improved specificity and coherence. In contrast, other methods show poorer tradeoffs: for example, the next-best method in Gemma is only 0.7% lower in

efficacy but suffers a 30% drop in similar-domain accuracy and an 8% drop in MMLU. Figure 4 reinforces these results, showing that PISCES outperforms the baselines by simultaneously attaining lower accuracy, and higher specificity and coherence scores. These results highlight that a precise, parameter-based approach to concept erasure enables finer-grained editing of model knowledge, yielding an improved efficacy-specificity tradeoff.

PISCES improves robustness to relearning Robustness evaluations in Table 1 reveal a substantial gap between PISCES and other methods. In Gemma, PISCES reaches a relearning accuracy of 51.5%, while the next-best method on efficacy reaches 85.4%—nearly 34% higher—indicating that most of the erased knowledge was recovered by fine-tuning on concept-related data, despite excluding evaluation answers. For Llama, PISCES performs slightly worse than in Gemma, reaching a relearning accuracy of 65.4%. However, other methods recover most or all of the removed knowledge, reaching 93.2%-103.1% accuracy post fine-tuning. This underscores that prior methods achieve only superficial concept erasure: the underlying knowledge remains in the model and can easily resurface. While PISCES also regains some knowledge under fine-tuning—leaving room for improvement—the up-to-38% gap in relearning accuracy shows that directly editing the parameters encoding the target concept yields substantially more robust erasure than general fine-tuning.

6 Analysis

To better understand the behavior and limitations of PISCES, we conduct two analyses. First, we study the relationship between the quality of the features identified by the disentangler and erasure success, highlighting the conditions under which PISCES performs best. Then, we compare the computational cost of PISCES to that of existing methods, showing that it offers a favorable trade-off between performance and efficiency.

6.1 Effect of Disentangler Performance on Erasure Success

A key component in our method is the disentangler model, which is used to identify concept-related features. Here, we analyze the relationship between the *quality* and *quantity* of features identified by the disentangler and the performance of PISCES.

⁴This results in a total of 800 experiments per concept for all methods and models.

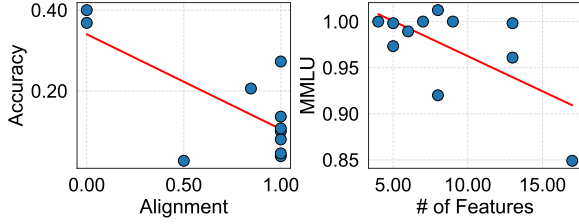


Figure 5: Analysis showing the relationships between feature alignment and erasure accuracy (left, -0.72 correlation with p -value 0.01), and between the number of selected features and MMLU performance (right, -0.64 correlation with p -value 0.03).

In our analysis, we consider the final set of selected features in Gemma-2-2B-IT.

To measure the quality of a feature f , we evaluate how well either the top-50 or bottom-50 tokens in its projection to the vocabulary (see Section 4.2) align with the target concept c . Let c' be our interpretation of the concept that f represents, we define two metrics:

1. *Alignment*: a binary score indicating whether c' aligns with c or not, i.e., 1 if c and c' are the same concepts and 0 otherwise. For example, a feature identified as relevant for the concept of $c = \text{baseball}$, but seems to represent the broader concept of $c' = \text{sports}$ will receive a score of 0.
2. *Coherence*: a discrete score from 0 to 2 which measures how clearly and distinctively c' is expressed among the top/bottom tokens in the projection, according to the presence of unrelated tokens. A score of 0 means *low coherence*, where no clear concept is observed. A score of 1 indicates *moderate coherence*, where f seems to encode c' but may also encode other concepts. A score of 2 indicates *high coherence*, where the tokens clearly reflect a single, well-defined concept aligned with c' .

Figure 5 presents the prominent patterns observed. Per-concept results and annotation examples can be found in §E. We find that features that strongly correspond to the target concept and express it clearly (i.e. high alignment and coherence) tend to yield better performance on our evaluation metrics. Moreover, concepts with many selected features often exhibit lower MMLU and Alpaca scores, likely due to accumulated reconstruction error (Gurnee, 2024). These results underscore that $\mathcal{D}$ PISCES relies on $\mathcal{D}$ ’s ability to identify precise, coherent features. When such features are present (e.g., “golf”, “Republic of Ireland”, “baseball”), $\mathcal{D}$ PISCES performs best; when they are ab-

Method	1 concept		10 concepts	
	Gemma	Llama	Gemma	Llama
MEMIT	$5 \cdot 10^{14}$	$1.9 \cdot 10^{15}$	$4.8 \cdot 10^{15}$	$5.8 \cdot 10^{16}$
AlphaEdit	$5.9 \cdot 10^{14}$	$2.4 \cdot 10^{15}$	$5.8 \cdot 10^{15}$	$2.3 \cdot 10^{16}$
ELM	$2.6 \cdot 10^{15}$	$1.1 \cdot 10^{16}$	$2.6 \cdot 10^{16}$	$1.1 \cdot 10^{17}$
RMU	$2.8 \cdot 10^{15}$	$1.1 \cdot 10^{16}$	$2.8 \cdot 10^{16}$	$1.1 \cdot 10^{17}$
PISCES	$5 \cdot 10^{14}$	$1.1 \cdot 10^{15}$	$5 \cdot 10^{14}$	$1.1 \cdot 10^{15}$

Table 2: Estimated FLOPs for applying each method to 1 and 10 concepts.

sent (e.g., “Uranium”), performance declines.

6.2 Computational Efficiency

In this section, we compare the computational cost of applying $\mathcal{D}$ PISCES versus other methods. We calculate the cost of $\mathcal{D}$ PISCES using $\mathcal{D}_{\text{SAE}}$ by summing the FLOPs to first perform vocabulary projection for every SAE feature vector, and then to apply the editing process for every isolated MLP vector. For RMU and ELM we rely on the heuristic FLOPs $\approx 6N$ for a forward and backward pass per token (Kaplan et al., 2020), multiplied by the amount of tokens in the forget and retain sets. Lastly, for MEMIT and AlphaEdit we approximate the cost by calculating the number of forward and backward passes needed for every fact in the forget set, and for calculating the covariance matrix and residual vector optimization.

Results are in Table 2, showing that $\mathcal{D}$ PISCES performs best at $5 \cdot 10^{14}$ FLOPs for Gemma, and $1.1 \cdot 10^{15}$ FLOPs for Llama, followed by MEMIT and AlphaEdit with similar performance, and then ELM and RMU which are one order of magnitude more expensive. Moreover, since running VocabProj can be performed once and reused across concepts, the cost of adding more concepts for $\mathcal{D}$ PISCES is comparatively insignificant. Therefore, when applying our method to multiple concepts, $\mathcal{D}$ PISCES becomes 1-2 orders of magnitude more efficient than all other methods. Notably, this analysis does not take into account the cost of training SAEs and assumes they are provided. Training a disentangler SAE is a preprocessing step for $\mathcal{D}$ PISCES, which can be done once rather than per concept. Yet, it entails a significant increase in the overall cost. To avoid this, one may consider alternative, more efficient disentanglers (see discussion in the Limitations section).

7 Conclusion

We present **🧩PISCES**, a framework for precisely erasing conceptual knowledge from language models by disentangling and directly editing their parameters. Unlike prior approaches that rely on fine-tuning or fact-level editing, **🧩PISCES** uses a disentangler model to isolate directions in the parameter space of the model that represent the concept and removes them with targeted edits. Experiments with two models and diverse concepts show that **🧩PISCES** achieves higher robustness and specificity than existing methods, while maintaining or slightly improving efficacy. These results establish in-parameter erasure as a state-of-the-art approach for fine-grained and robust conceptual knowledge removal in LLMs.

Limitations

Although **🧩PISCES** performs well in our evaluations, there remains significant room for improvement. First, our current implementation only targets the MLP parameters. While prior work has shown that MLPs encode knowledge in the model (Geva et al., 2021, 2022; Dai et al., 2022; Meng et al., 2022; Geva et al., 2023), recent findings suggest that attention heads also contribute to knowledge storage (Elhelo and Geva, 2024). Extending **🧩PISCES** to include these components could enable more comprehensive erasure.

Second, our reliance on SAEs for the disentangler introduces limitations. We can only erase concepts that were captured as features, and must contend with imperfect reconstructions. Future work establishing new methods for disentangling model parameters could address these limitations, and thanks to the generality of **🧩PISCES**, be easily integrated into our framework. Another possible direction could be to explore supervised disentanglement approaches (Geiger et al., 2024; Huang et al., 2024) as potential alternatives to the current unsupervised setup—a possibility we leave for future investigation.

Lastly, we identify concept-related features based on VocabProj. While this method has proven effective for identifying causal effects on model outputs, it is less reliable in early layers. Thus, incorporating complementary automated interpretability techniques for identifying concept-related features could potentially improve the overall performance.

Ethical Considerations

Our work introduces **🧩PISCES**, a framework for precise in-parameter erasure of conceptual knowledge in language models. While the goal is to enable removal of undesirable or sensitive concepts, such as fictional content or protected information, this capability could in principle be misused for censorship or the suppression of legitimate knowledge. We acknowledge this risk, but believe the potential benefits of our method outweigh it: enabling safer deployment of LLMs by removing inappropriate or restricted content, supporting compliance with copyright obligations, and enabling better understanding of how concepts are encoded in model parameters. We hope that the insights and tools provided in this work are used to support responsible and transparent AI development.

Acknowledgments

This work was supported in part by the Gemma 2 Academic Research Program at Google, the Alon scholarship, and the Israel Science Foundation grant 1083/24. Figures 2 and 3 use images from www.freepik.com.

References

- Andy Arditi, Oscar Obeso, Aaquib Syed, Daniel Paleka, Nina Panickssery, Wes Gurnee, and Neel Nanda. 2024. Refusal in language models is mediated by a single direction. *Advances in Neural Information Processing Systems*, 37:136037–136083.
- Tomer Ashuach, Martin Tutek, and Yonatan Belinkov. 2025. **REVS: Unlearning sensitive information in language models via rank editing in the vocabulary space**. In *Findings of the Association for Computational Linguistics: ACL 2025*, pages 14774–14797, Vienna, Austria. Association for Computational Linguistics.
- Fazl Barez, Tingchen Fu, Ameya Prabhu, Stephen Casper, Amartya Sanyal, Adel Bibi, Aidan O’Gara, Robert Kirk, Ben Bucknall, Tim Fist, Luke Ong, Philip H. S. Torr, Kwok-Yan Lam, Robert F. Trager, David Krueger, Sören Mindermann, José Hernández-Orallo, Mor Geva, and Yarin Gal. 2025. **Open problems in machine unlearning for ai safety**. *ArXiv*, abs/2501.04952.
- Nora Belrose, David Schneider-Joseph, Shauli Ravfogel, Ryan Cotterell, Edward Raff, and Stella Biderman. 2023. **LEACE: Perfect linear concept erasure in closed form**. In *Thirty-seventh Conference on Neural Information Processing Systems*.
- Tolga Bolukbasi, Kai-Wei Chang, James Y. Zou, Venkatesh Saligrama, and Adam Tauman Kalai. 2016.

- Man is to computer programmer as woman is to homemaker? debiasing word embeddings. In *Neural Information Processing Systems*.
- Trenton Bricken, Adly Templeton, Joshua Batson, Brian Chen, Adam Jermy, Tom Conerly, Nick Turner, Cem Anil, Carson Denison, Amanda Askell, Robert Lasenby, Yifan Wu, Shauna Kravec, Nicholas Schiefer, Tim Maxwell, Nicholas Joseph, Zac Hatfield-Dodds, Alex Tamkin, Karina Nguyen, and 6 others. 2023. Towards monosemanticity: Decomposing language models with dictionary learning. *Transformer Circuits Thread*. <https://transformer-circuits.pub/2023/monosemantic-features/index.html>.
- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel Ziegler, Jeffrey Wu, Clemens Winter, and 12 others. 2020. [Language models are few-shot learners](#). In *Advances in Neural Information Processing Systems*, volume 33, pages 1877–1901. Curran Associates, Inc.
- Nitay Calderon, Roi Reichart, and Rotem Dror. 2025. [The alternative annotator test for llm-as-a-judge: How to statistically justify replacing human annotators with llms](#). *Preprint*, arXiv:2501.10970.
- Yinzhi Cao and Junfeng Yang. 2015. [Towards making systems forget with machine unlearning](#). 2015 *IEEE Symposium on Security and Privacy*, pages 463–480.
- Yuheng Chen, Pengfei Cao, Kang Liu, and Jun Zhao. 2025. [The knowledge microscope: Features as better analytical lenses than neurons](#). *Preprint*, arXiv:2502.12483.
- Damai Dai, Li Dong, Yaru Hao, Zhifang Sui, Baobao Chang, and Furu Wei. 2022. [Knowledge neurons in pretrained transformers](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8493–8502, Dublin, Ireland. Association for Computational Linguistics.
- Fahim Dalvi, Abdul Rafae Khan, Firoj Alam, Nadir Durani, Jia Xu, and Hassan Sajjad. 2022. [Discovering latent concepts learned in BERT](#). In *International Conference on Learning Representations*.
- Aghyad Deeb and Fabien Roger. 2025. [Do unlearning methods remove information from language model weights?](#)
- Jai Doshi and Asa Cooper Stickland. 2025. [Does unlearning truly unlearn? a black box evaluation of llm unlearning methods](#). *Preprint*, arXiv:2411.12103.
- Abhimanyu Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, Anirudh Goyal, Anthony Hartshorn, Aobo Yang, Archi Mitra, Archie Sravankumar, Artem Korenev, Arthur Hinsvark, Arun Rao, Aston Zhang, and 82 others. 2024. [The llama 3 herd of models](#). *CoRR*, abs/2407.21783.
- Ronen Eldan and Mark Russinovich. 2023. [Who’s harry potter? approximate unlearning in llms](#). *Preprint*, arXiv:2310.02238.
- Nelson Elhage, Tristan Hume, Catherine Olsson, Nicholas Schiefer, Tom Henighan, Shauna Kravec, Zac Hatfield-Dodds, Robert Lasenby, Dawn Drain, Carol Chen, Roger Grosse, Sam McCandlish, Jared Kaplan, Dario Amodei, Martin Wattenberg, and Christopher Olah. 2022. Toy models of superposition. *Transformer Circuits Thread*.
- Amit Elhelo and Mor Geva. 2024. Inferring functionality of attention heads from their parameters. *arXiv preprint arXiv:2412.11965*.
- Junfeng Fang, Houcheng Jiang, Kun Wang, Yunshan Ma, Jie Shi, Xiang Wang, Xiangnan He, and Tat-Seng Chua. 2025. [Alphaedit: Null-space constrained model editing for language models](#). In *The Thirteenth International Conference on Learning Representations*.
- Eoin Farrell, Yeu-Tong Lau, and Arthur Conmy. 2024. [Applying sparse autoencoders to unlearn knowledge in language models](#). *Preprint*, arXiv:2410.19278.
- Joseph L Fleiss. 1971. Measuring nominal scale agreement among many raters. *Psychological bulletin*, 76(5):378.
- Ahmed Frikha, Muhammad Reza Ar Razi, Krishna Kanth Nakka, Ricardo Mendes, Xue Jiang, and Xuebing Zhou. 2025. [Privacyscalpel: Enhancing llm privacy via interpretable feature intervention with sparse autoencoders](#). *Preprint*, arXiv:2503.11232.
- Rohit Gandikota, Sheridan Feucht, Samuel Marks, and David Bau. 2025. [Erasing conceptual knowledge from language models](#).
- Leo Gao, Tom Dupre la Tour, Henk Tillman, Gabriel Goh, Rajan Troll, Alec Radford, Ilya Sutskever, Jan Leike, and Jeffrey Wu. 2025. [Scaling and evaluating sparse autoencoders](#). In *The Thirteenth International Conference on Learning Representations*.
- Atticus Geiger, Zhengxuan Wu, Christopher Potts, Thomas Icard, and Noah D. Goodman. 2024. [Finding alignments between interpretable causal variables and distributed neural representations](#). In *Causal Learning and Reasoning, 1-3 April 2024, Los Angeles, California, USA*, volume 236 of *Proceedings of Machine Learning Research*, pages 160–187. PMLR.
- Mor Geva, Jasmijn Bastings, Katja Filippova, and Amir Globerson. 2023. [Dissecting recall of factual associations in auto-regressive language models](#). In *The 2023 Conference on Empirical Methods in Natural Language Processing*.

- Mor Geva, Avi Caciularu, Kevin Wang, and Yoav Goldberg. 2022. [Transformer feed-forward layers build predictions by promoting concepts in the vocabulary space](#). In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 30–45, Abu Dhabi, United Arab Emirates. Association for Computational Linguistics.
- Mor Geva, Roei Schuster, Jonathan Berant, and Omer Levy. 2021. [Transformer feed-forward layers are key-value memories](#). In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing*, pages 5484–5495, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Amirata Ghorbani, James Wexler, James Y. Zou, and Been Kim. 2019. [Towards automatic concept-based explanations](#). In *Neural Information Processing Systems*.
- Antonio Ginart, Melody Guan, Gregory Valiant, and James Y Zou. 2019. Making ai forget you: Data deletion in machine learning. *Advances in neural information processing systems*, 32.
- Yichen Gong, Delong Ran, Xinlei He, Tianshuo Cong, Anyu Wang, and Xiaoyun Wang. 2025. [Safety misalignment against large language models](#). *Proceedings 2025 Network and Distributed System Security Symposium*.
- Google. 2025. Gemini 2.0 Flash. <https://cloud.google.com/vertex-ai/generative-ai/docs/models/gemini/2-0-flash>.
- Kathrin Grosse, Lukas Bieringer, Tarek R. Besold, and Alexandre Alahi. 2024. Towards more practical threat models in artificial intelligence security. In *Proceedings of the 33rd USENIX Conference on Security Symposium*, SEC ’24, USA. USENIX Association.
- Yoav Gur-Arieh, Roy Mayan, Chen Agassy, Atticus Geiger, and Mor Geva. 2025. Enhancing automated interpretability with output-centric feature descriptions. In *The 63rd Annual Meeting of the Association for Computational Linguistics*.
- Wes Gurnee. 2024. Sae reconstruction errors are (empirically) pathological. In *AI Alignment Forum*, page 16.
- Wes Gurnee, Neel Nanda, Matthew Pauly, Katherine Harvey, Dmitrii Troitskii, and Dimitris Bertsimas. 2023. [Finding neurons in a haystack: Case studies with sparse probing](#). *Transactions on Machine Learning Research*.
- Zhengfu He, Wentao Shu, Xuyang Ge, Lingjie Chen, Junxuan Wang, Yunhua Zhou, Frances Liu, Qipeng Guo, Xuanjing Huang, Zuxuan Wu, Yu-Gang Jiang, and Xipeng Qiu. 2024. [Llama scope: Extracting millions of features from llama-3.1-8b with sparse autoencoders](#). *ArXiv*, abs/2410.20526.
- Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. 2021. [Measuring massive multitask language understanding](#). In *International Conference on Learning Representations*.
- Yihuai Hong, Lei Yu, Haiqin Yang, Shauli Ravfogel, and Mor Geva. 2025. [Intrinsic evaluation of unlearning using parametric knowledge traces](#).
- Cheng-Hsun Hsueh, Paul Kuo-Ming Huang, Tzu-Han Lin, Che Wei Liao, Hung-Chieh Fang, Chao-Wei Huang, and Yun-Nung Chen. 2024. [Editing the mind of giants: An in-depth exploration of pitfalls of knowledge editing in large language models](#). In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 9417–9429, Miami, Florida, USA. Association for Computational Linguistics.
- Chenhui Hu, Pengfei Cao, Yubo Chen, Kang Liu, and Jun Zhao. 2024. [WilKE: Wise-layer knowledge editor for lifelong knowledge editing](#). In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 3476–3503, Bangkok, Thailand. Association for Computational Linguistics.
- Shengyuan Hu, Yiwei Fu, Steven Wu, and Virginia Smith. 2025. [Unlearning or obfuscating? jogging the memory of unlearned LLMs via benign relearning](#). In *The Thirteenth International Conference on Learning Representations*.
- Jing Huang, Zhengxuan Wu, Christopher Potts, Mor Geva, and Atticus Geiger. 2024. [RAVEL: Evaluating interpretability methods on disentangling language model representations](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 8669–8687, Bangkok, Thailand. Association for Computational Linguistics.
- Lei Huang, Weijiang Yu, Weitao Ma, Weihong Zhong, Zhangyin Feng, Haotian Wang, Qianglong Chen, Weihua Peng, Xiaocheng Feng, Bing Qin, and Ting Liu. 2025. [A survey on hallucination in large language models: Principles, taxonomy, challenges, and open questions](#). *ACM Trans. Inf. Syst.*, 43(2).
- Robert Huben, Hoagy Cunningham, Logan Riggs Smith, Aidan Ewart, and Lee Sharkey. 2024. [Sparse autoencoders find highly interpretable features in language models](#). In *The Twelfth International Conference on Learning Representations*.
- OpenAI Aaron Hurst, Adam Lerer, Adam P. Goucher, Adam Perelman, Aditya Ramesh, Aidan Clark, AJ Ostrow, Akila Welihinda, Alan Hayes, Alec Radford, Aleksander Mkadry, Alex Baker-Whitcomb, Alex Beutel, Alex Borzunov, Alex Carney, Alex Chow, Alexander Kirillov, Alex Nichol, Alex Paino, and 397 others. 2024. [Gpt-4o system card](#). *ArXiv*, abs/2410.21276.
- Shadi Iskander, Kira Radinsky, and Yonatan Belinkov. 2023. [Shielded representations: Protecting sensitive](#)

- attributes through iterative gradient-based projection. In *Annual Meeting of the Association for Computational Linguistics*.
- Curt Tigges Joseph Bloom and David Chanin. 2024. Saelens. <https://github.com/jbloomAus/SAELens>.
- Jared Kaplan, Sam McCandlish, Tom Henighan, Tom B. Brown, Benjamin Chess, Rewon Child, Scott Gray, Alec Radford, Jeffrey Wu, and Dario Amodei. 2020. *Scaling laws for neural language models*. Preprint, arXiv:2001.08361.
- Yassine El Kheir, Ahmed Ali, and Shammur A. Chowdhury. 2024. *Speech representation analysis based on inter- and intra-model similarities*. 2024 *IEEE International Conference on Acoustics, Speech, and Signal Processing Workshops (ICASSPW)*, pages 848–852.
- Connor Kissane, Robert Krzyzanowski, Joseph Isaac Bloom, Arthur Conmy, and Neel Nanda. 2024. *Interpreting attention layer outputs with sparse autoencoders*. Preprint, arXiv:2406.17759.
- J. Richard Landis and Gary G. Koch. 1977. *The measurement of observer agreement for categorical data*. *Biometrics*, 33(1):159–174.
- Quoc V. Le, Marc’Aurelio Ranzato, Rajat Monga, Matthieu Devin, Gregory S. Corrado, Kai Chen, Jeffrey Dean, and A. Ng. 2011. *Building high-level features using large scale unsupervised learning*. 2013 *IEEE International Conference on Acoustics, Speech and Signal Processing*, pages 8595–8598.
- Honglak Lee, Chaitanya Ekanadham, and A. Ng. 2007. *Sparse deep belief net model for visual area v2*. In *Neural Information Processing Systems*.
- Quentin Lhoest, Albert Villanova del Moral, Yacine Jernite, Abhishek Thakur, Patrick von Platen, Suraj Patil, Julien Chaumond, Mariama Drame, Julien Plu, Lewis Tunstall, Joe Davison, Mario Šaško, Gunjan Chhablani, Bhavitvya Malik, Simon Brandeis, Teven Le Scao, Victor Sanh, Canwen Xu, Nicolas Patry, and 13 others. 2021. *Datasets: A community library for natural language processing*. In *Proceedings of the 2021 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*, pages 175–184, Online and Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Nathaniel Li, Alexander Pan, Anjali Gopal, Summer Yue, Daniel Berrios, Alice Gatti, Justin D. Li, Ann-Kathrin Dombrowski, Shashwat Goel, Gabriel Mukobi, Nathan Helm-Burger, Rassin Lababidi, Lennart Justen, Andrew Bo Liu, Michael Chen, Isabelle Barrass, Oliver Zhang, Xiaoyuan Zhu, Rishub Tamirisa, and 27 others. 2024a. *The WMDP benchmark: Measuring and reducing malicious use with unlearning*. In *Forty-first International Conference on Machine Learning*.
- Xuechen Li, Tianyi Zhang, Yann Dubois, Rohan Taori, Ishaan Gulrajani, Carlos Guestrin, Percy Liang, and Tatsunori B. Hashimoto. 2023. AlpacaEval: An Automatic Evaluator of Instruction-following Models.
- Yanhong Li, Chunling Fan, Mingqing Huang, and Chengming Li. 2024b. *Learning from mistakes: A comprehensive review of knowledge editing for large language models*. In *2024 IEEE International Conference on Smart Internet of Things (SmartIoT)*, pages 563–569.
- Tom Lieberum, Senthoooran Rajamanoharan, Arthur Conmy, Lewis Smith, Nicolas Sonnerat, Vikrant Varma, Janos Kramar, Anca Dragan, Rohin Shah, and Neel Nanda. 2024. *Gemma scope: Open sparse autoencoders everywhere all at once on gemma 2*. In *The 7th BlackboxNLP Workshop*.
- Hanxiao Liu, Zihang Dai, David So, and Quoc V Le. 2021. *Pay attention to MLPs*. In *Advances in Neural Information Processing Systems*.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. *Roberta: A robustly optimized bert pretraining approach*. arXiv preprint arXiv:1907.11692.
- Zheyuan Liu, Guangyao Dou, Zhaoxuan Tan, Yijun Tian, and Meng Jiang. 2024. *Machine unlearning in generative ai: A survey*. ArXiv, abs/2407.20516.
- Ziyao Liu, Huanyi Ye, Chen Chen, Yongsen Zheng, and Kwok-Yan Lam. 2025. *Threats, attacks, and defenses in machine unlearning: A survey*. *IEEE Open Journal of the Computer Society*.
- Michelle Lo, Shay B. Cohen, and Fazl Barez. 2024. *Large language models relearn removed concepts*. Preprint, arXiv:2401.01814.
- Jakub Łucki, Boyi Wei, Yangsibo Huang, Peter Henderson, Florian Tramèr, and Javier Rando. 2025. *An adversarial perspective on machine unlearning for AI safety*. *Transactions on Machine Learning Research*.
- Aengus Lynch, Phillip Guo, Aidan Ewart, Stephen Casper, and Dylan Hadfield-Menell. 2024. *Eight methods to evaluate robust unlearning in llms*. Preprint, arXiv:2402.16835.
- Samuel Marks, Can Rager, Eric J Michaud, Yonatan Belinkov, David Bau, and Aaron Mueller. 2025. *Sparse feature circuits: Discovering and editing interpretable causal graphs in language models*. In *The Thirteenth International Conference on Learning Representations*.
- Kevin Meng, David Bau, Alex Andonian, and Yonatan Belinkov. 2022. *Locating and editing factual associations in gpt*. In *Neural Information Processing Systems*.

- Kevin Meng, Arnab Sen Sharma, Alex J Andonian, Yonatan Belinkov, and David Bau. 2023. [Mass-editing memory in a transformer](#). In *The Eleventh International Conference on Learning Representations*.
- Eric Mitchell, Charles Lin, Antoine Bosselut, Chelsea Finn, and Christopher D Manning. 2022. [Fast model editing at scale](#). In *International Conference on Learning Representations*.
- Aashiq Muhamed, Jacopo Bonato, Mona Diab, and Virginia Smith. 2025. Saes can improve unlearning: Dynamic sparse autoencoder guardrails for precision unlearning in llms. *arXiv preprint arXiv:2504.08192*.
- Neel Nanda and Joseph Bloom. 2022. Transformerlens. <https://github.com/TransformerLensOrg/TransformerLens>.
- Nostalgebraist. 2020. [interpreting GPT: the logit lens](#).
- OpenAI. 2025. [Openai o3 and o4-mini system card](#).
- Fabio Petroni, Tim Rocktäschel, Sebastian Riedel, Patrick Lewis, Anton Bakhtin, Yuxiang Wu, and Alexander Miller. 2019. [Language models as knowledge bases?](#) In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 2463–2473, Hong Kong, China. Association for Computational Linguistics.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. [Language models are unsupervised multitask learners](#). *OpenAI*. Accessed: 2024-11-15.
- Pranav Rajpurkar, Robin Jia, and Percy Liang. 2018. [Know what you don’t know: Unanswerable questions for SQuAD](#). In *Proceedings of the 56th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, pages 784–789, Melbourne, Australia. Association for Computational Linguistics.
- Juan Ramos and 1 others. 2003. Using tf-idf to determine word relevance in document queries. In *Proceedings of the first instructional conference on machine learning*, volume 242, pages 29–48. Citeseer.
- Shauli Ravfogel, Yanai Elazar, Hila Gonen, Michael Twiton, and Yoav Goldberg. 2020. [Null it out: Guarding protected attributes by iterative nullspace projection](#). In *Annual Meeting of the Association for Computational Linguistics*.
- Nils Reimers and Iryna Gurevych. 2019. [Sentence-bert: Sentence embeddings using siamese bert-networks](#). In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing*. Association for Computational Linguistics.
- Nina Rimskey, Nick Gabrieli, Julian Schulz, Meg Tong, Evan Hubinger, and Alexander Turner. 2024. [Steering llama 2 via contrastive activation addition](#). In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 15504–15522, Bangkok, Thailand. Association for Computational Linguistics.
- Gemma Team Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, L’eonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ram’e, Johan Ferret, Peter Liu, Pouya Dehghani Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, Anton Tsitsulin, and 176 others. 2024. [Gemma 2: Improving open language models at a practical size](#). *ArXiv*, abs/2408.00118.
- Adam Roberts, Colin Raffel, and Noam Shazeer. 2020. [How much knowledge can you pack into the parameters of a language model?](#) In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 5418–5426, Online. Association for Computational Linguistics.
- Hassan Sajjad, Nadir Durrani, and Fahim Dalvi. 2021. [Neuron-level interpretation of deep nlp models: A survey](#). *Transactions of the Association for Computational Linguistics*, 10:1285–1303.
- Hassan Sajjad, Nadir Durrani, Fahim Dalvi, Firoj Alam, Abdul Khan, and Jia Xu. 2022. [Analyzing encoded concepts in transformer language models](#). In *Proceedings of the 2022 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, pages 3082–3101, Seattle, United States. Association for Computational Linguistics.
- Pratiksha Thaker, Shengyuan Hu, Neil Kale, Yash Maurya, Zhiwei Steven Wu, and Virginia Smith. 2024. [Position: Llm unlearning benchmarks are weak measures of progress](#). *ArXiv*, abs/2410.02879.
- Elena Voita, Javier Ferrando, and Christoforos Nalmpantis. 2024. [Neurons in large language models: Dead, n-gram, positional](#). In *Findings of the Association for Computational Linguistics: ACL 2024*, pages 1288–1301, Bangkok, Thailand. Association for Computational Linguistics.
- T Wolf. 2019. Huggingface’s transformers: State-of-the-art natural language processing. *arXiv preprint arXiv:1910.03771*.
- Ruihan Wu, Chhavi Yadav, Russ Salakhutdinov, and Kamalika Chaudhuri. 2025a. [Evaluating deep unlearning in large language models](#).
- Xinwei Wu, Junzhuo Li, Minghui Xu, Weilong Dong, Shuangzhi Wu, Chao Bian, and Deyi Xiong. 2023. [Depn: Detecting and editing privacy neurons in pre-trained language models](#). In *Conference on Empirical Methods in Natural Language Processing*.

Zhengxuan Wu, Aryaman Arora, Atticus Geiger, Zheng Wang, Jing Huang, Dan Jurafsky, Christopher D. Manning, and Christopher Potts. 2025b. [Axbench: Steering llms? even simple baselines outperform sparse autoencoders](#). *Preprint*, arXiv:2501.17148.

Zhengxuan Wu, Aryaman Arora, Atticus Geiger, Zheng Wang, Jing Huang, Dan Jurafsky, Christopher D. Manning, and Christopher Potts. 2025c. Axbench: Steering llms? even simple baselines outperform sparse autoencoders. *arXiv preprint arXiv:2501.17148*.

Tomoya Yamashita, Takayuki Miura, Yuuki Yamanaka, Toshiaki Shibahara, and Masanori Yamada. 2024. [Concept unlearning for large language models](#). In *Neurips Safe Generative AI Workshop 2024*.

Ruiqi Zhang, Licong Lin, Yu Bai, and Song Mei. 2024. [Negative preference optimization: From catastrophic collapse to effective unlearning](#). In *First Conference on Language Modeling*.

Andy Zou, Long Phan, Justin Wang, Derek Duenas, Maxwell Lin, Maksym Andriushchenko, J Zico Kolter, Matt Fredrikson, and Dan Hendrycks. 2024. [Improving alignment and robustness with circuit breakers](#). In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*.

Andy Zou, Zifan Wang, J. Zico Kolter, and Matt Fredrikson. 2023. [Universal and transferable adversarial attacks on aligned language models](#). *ArXiv*, abs/2307.15043.

A Method Implementation Details

A.1 SAE Feature Selection

To select features that are relevant to a given concept, we first identify tokens associated with that concept. This section outlines the process we followed, using the “*Culture of Greece*” concept and Gemma-2-2B-IT model as a running example (Riviere et al., 2024).

Token Selection. We begin by constructing a concept-specific token set:

1. We tokenize the forget set associated with the target concept, removing stop words to reduce noise.
2. We apply a TF-IDF model (Ramos et al., 2003) to identify the most informative tokens in the filtered text.
3. We manually select 2–5 tokens that appear highly correlated with the concept, preferably from among the top TF-IDF tokens.

Example: For the “*Culture of Greece*” concept, we selected ‘Greek’, ‘Greece’,

and ‘Athens’. TF-IDF ranked ‘Greek’ and ‘Greece’ as the top two tokens, with ‘Athens’ ranked 11th.

4. We automatically expand the manually selected set by:

- Including tokens that match the selected ones, ignoring case.
- Adding tokens that are similar in the model’s embedding space (measured by cosine similarity).

Example: Expanding the selected tokens led to the following set: (‘greece’, ‘Athens’, ‘Athens’, ‘greek’, ‘GREEK’, ‘Greece’, ‘greek’, ‘Greek’, ‘Greeks’, ‘Greece’, ‘Athenian’, ‘Griechenland’, ‘Greek’, ‘griech’).

Feature Selection. Using the final token set, we then identify and filter relevant SAE features:

1. For each SAE feature, we apply VocabProj to obtain the tokens most associated with it.
2. We compute the intersection between the associated tokens and the token set. Features with an intersection size greater than a threshold α (we used $\alpha = 4$) are selected.
3. From this candidate set, we manually filter features that appear strongly aligned with the target concept and weakly associated with unrelated concepts. This manual step typically takes under a minute.

Example:

- We retained feature [‘Greek’, ‘Greek’, ‘GREEK’, ‘greek’, ‘Greeks’, ‘Greece’, ‘griech’, ‘grecque’].
 - We rejected feature [‘Italians’, ‘austria’, ‘Americans’, ‘Spaniards’, ‘Egyptians’, ‘Tajikistan’, ‘Greece’, ‘Americans’] due to its overlap with unrelated concepts.
4. Finally, we prune features by measuring their individual impact on model behavior under our editing procedure. Any feature whose ablation leads to a significant performance degradation, as measured on the MMLU validation set, is discarded.

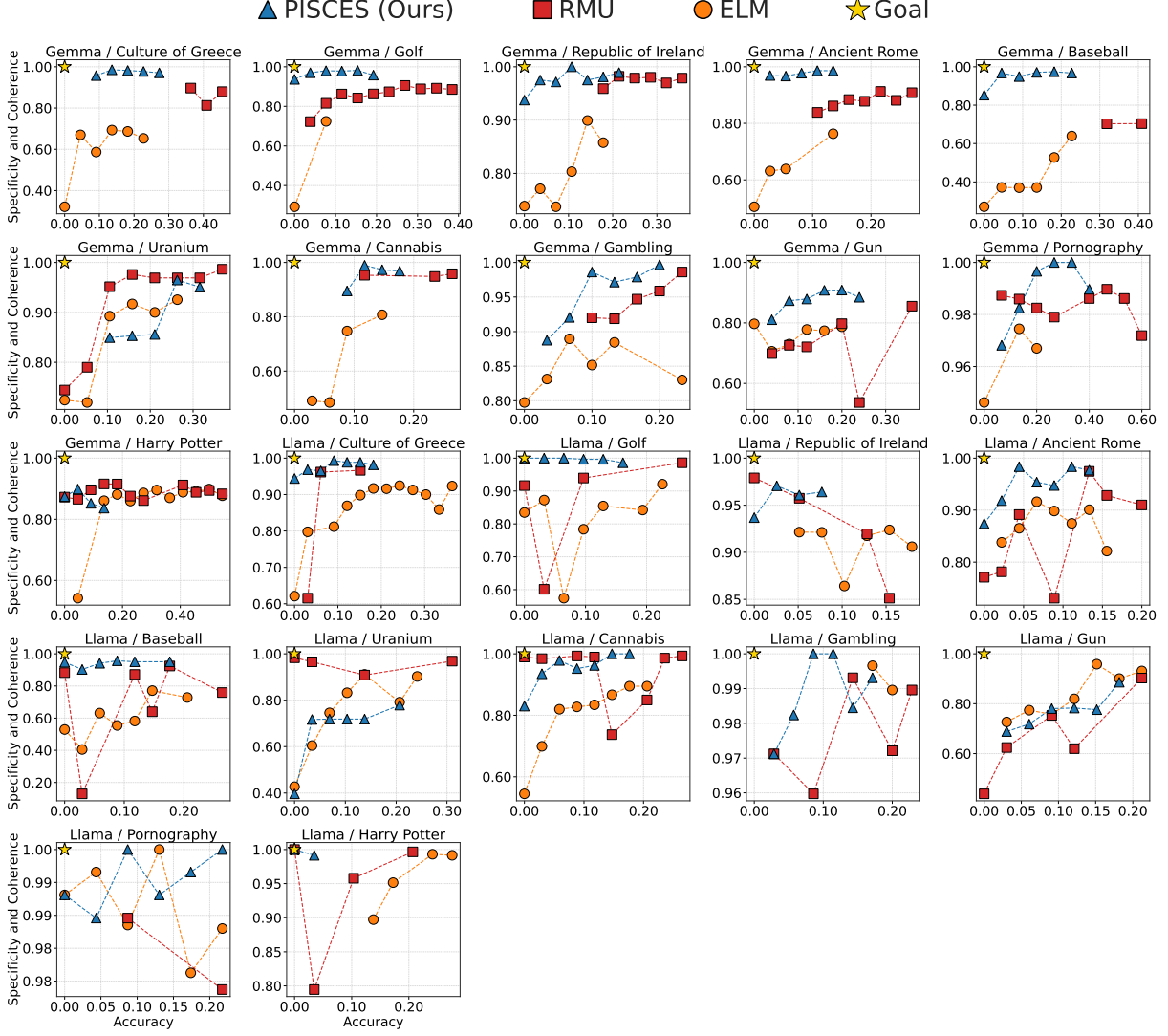


Figure 6: Performance of PISCES, ELM and RMU on all concepts and two models (Gemma-2-2b-it and Llama-3.1-8b-it). Each point is a single hyperparameter selection taken out of 100 possible choices, presenting only the best performing ones. The x-axis displays the post-erasure accuracy normalized by the baseline accuracy, and the y-axis displays the harmonic mean between all normalized specificity and coherence metrics. The star represents the goal – zero accuracy and 100% specificity and coherence.

A.2 Setting SAE Feature Activations

When editing MLP vectors using SAEs by disentangling them and affecting specific features’ activations, we must take care to affect them correctly such that we don’t cause the opposite effect to the one we were pursuing. This is because an MLP vector that seems to promote a concept c , might actually be used by the model to suppress it, through negative activations. Therefore, using the notations from §4.1 where $\mathbf{v}_i$ is an MLP vector we’re editing, a_i is its activation, and f is a targeted feature, we must identify two factors: (1) Does f promote or suppress c , and (2) is a_i positive or negative in the concept’s context. We determine

(1) by whether concept-related tokens appear in the top of the feature vector’s vocabulary projection, or the bottom (Voita et al., 2024). We can then ascertain (2) by feeding the concept’s forget-set data through the model and taking the majority sign of a_i . We then set s_f to 1 (−1) if f promotes (suppresses) c , and s_{a_i} to be a_i ’s majority sign as described above. Finally, when editing $\mathbf{v}_i$ we set $\bar{\mathbf{m}}_f^i = -(s_f \cdot s_{a_i}) \cdot \mu \cdot \hat{\mathbf{m}}_f$.

A.3 Evaluating Feature Selection Agreement

Since our proposed method requires a brief manual feature filtering stage, we conduct a human evaluation assessing agreement between annotators. For each model, we randomly sampled 5 concepts and

compiled their respective feature candidate sets, resulting in a total of 10 concepts and 158 features. We then assigned four annotators (NLP graduate students) to decide whether to include or exclude each of the candidate features of each of the ten concept. Across all candidate features, inter-annotator agreement measured by Fleiss’ κ was 0.574, indicating moderate to near substantial agreement (Fleiss, 1971; Landis and Koch, 1977).

B Hyperparameter Selection

To attain the best possible performance per concept, we conduct a hyperparameter grid search for each method per concept. We define 100 hyperparameter configurations based on prior work and manual tuning informed by the original papers. Each method is evaluated on a validation set disjoint from the test set, and we select the configuration that achieves the highest harmonic mean of efficacy, specificity, and coherence.

For **PISCES** we selected the range $\mu \in \{4, 7, 10, 13, 18, 24, 30, 36, 42, 50\}$ and $\tau \in \{0.2, 0.3, 0.4, 0.5, 0.7, 0.75, 0.8, 0.85, 0.9, 0.95\}$, allowing for a broad range of activation strengths and widths. For **ELM**, we selected $\eta \in \{1000, 2000, 5000\}$, $\alpha \in \{8, 16, 32\}$ and 11 numbers of epochs evenly distributed between 40 and 440 – including the latter as we saw that it made a significant difference in the method’s efficacy-specificity tradeoff. For **RMU** we selected steering coefficient $\in \{3, 6, 9, 12, 15, 18, 21, 24, 27, 30\}$ and $\alpha \in \{3, 5, 8, 12, 25, 50, 100, 200, 300, 600\}$. For **MEMIT** in Llama we focus the edit on layers 4,5,6,7,8, with learning rates, optimization steps, and clamp norm factors in the ranges $[1 \cdot 10^{-1}, 2 \cdot 10^{-1}, 3 \cdot 10^{-1}, 4 \cdot 10^{-1}, 5 \cdot 10^{-1}]$, $[10, 15, 20, 25, 30]$ and $[1, 2, 5, 7, 10, 14, 20]$ respectively. In Gemma we focus on the edit layer 3,4,5,6,7, with learning rates and optimization steps and clamp norm factors in the ranges $[1 \cdot 10^{-1}, 3 \cdot 10^{-1}, 5 \cdot 10^{-1}]$, $[5, 10, 20]$ and $\{0.5, 0.75, 1, 2, 4, 5, 7, 9, 11, 13, 15\}$ respectively. Finally, for AlphaEdit in Llama we focus on the same layers with clamp norm factors, learning rates, and optimization steps of $\{2, 4, 6, 8, 12, 16, 24, 40\}$, $\{0.1, 0.3, 0.5\}$ and $\{20, 25, 30, 35\}$ respectively. For Gemma we had $\{0.75, 1, 2, 4, 8, 16\}$, $\{0.1, 0.2, 0.3, 0.5\}$ and $\{5, 10, 15, 20, 25\}$ respectively. To perform MEMIT and AlphaEdit we follow the steps in

(Hong et al., 2025).

C Justifying use of LLM-as-a-Judge

To justify our use of an LLM-as-a-Judge for evaluating model-generated answers, we apply the alternative annotator test proposed by Calderon et al. (2025), which assesses whether the LLM performs as well as or better than a randomly selected human annotator. Following their procedure, we recruited three human annotators (graduate students) and used a set of 120 questions sampled uniformly across concepts, methods, models, and accuracy-based evaluations. For each question, annotators received the same inputs as the LLM judge: the question, the correct answer, and the model’s generated answer. They were asked to evaluate whether the model’s answer matched the correct one (instructions can be seen in Figure 9). Following Calderon et al. (2025), we set $\epsilon = 0.1$ to reflect the low-expertise nature of the task. The analysis yielded a winning rate of $\omega = 0.67$ with a p-value of 0.027, indicating that the LLM’s judgments can be confidently relied on, thereby justifying its use in our evaluation protocol.

D Data Generation

D.1 Generating Questions

To generate questions for measuring accuracy and similar domain accuracy, we use the GPT-o3 model (OpenAI, 2025). The following are the prompts used for generating the questions.

Question Generation Prompt

I’m going to give you each time a concept and its Wikipedia entry. Your job is, based on this Wikipedia knowledge and your general knowledge about the concept, to:

- Generate exactly 100 questions and answer in JSON format (question under “q” key and answer under “a”).
- The questions should be easy and simple.
- They should be answerable from knowledge of the concept—no current or changing information.
- They must not repeat themselves.
- Each should have **one** specific correct answer (avoid “Name one X”-style prompts).
- Questions must concern the concept only, not tangential topics.
- Ensure every question has exactly one specific correct answer.

Concept: {concept}

Wikipedia entry: {wikipedia_text}

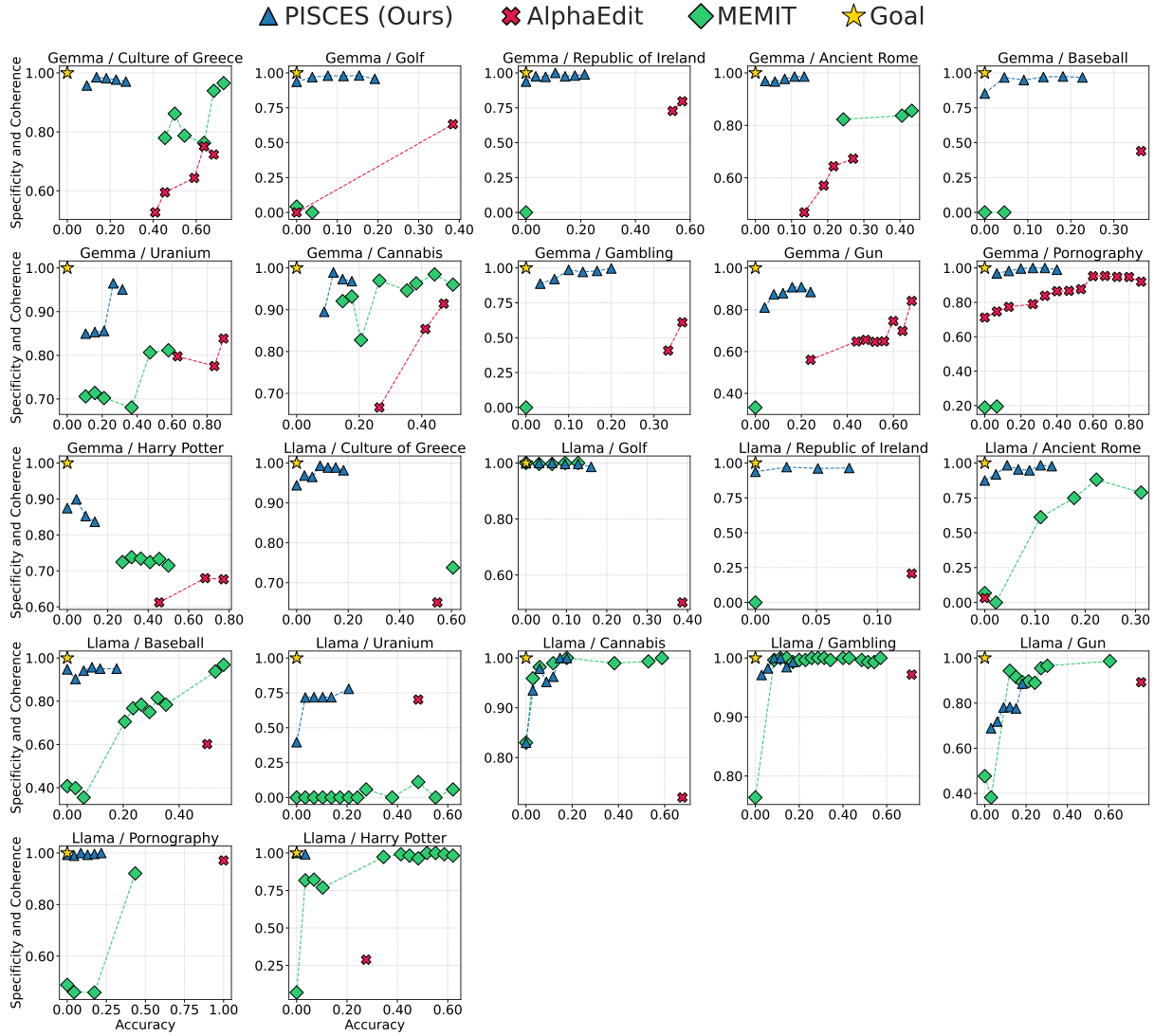


Figure 7: Performance of PISCES, MEMIT and AlphaEdit on all concepts and two models (Gemma-2-2b-it and Llama-3.1-8b-it). Each point is a single hyperparameter selection taken out of 100 possible choices, presenting only the best performing ones. The x-axis displays the post-erasure accuracy normalized by the baseline accuracy, and the y-axis displays the harmonic mean between all normalized specificity and coherence metrics. The star represents the goal – zero accuracy and 100% specificity and coherence.

Similar Domain Prompt

Each time I'm going to give you a concept. Please then generate 100 questions and answers in a JSON format, with questions under "q" and answers under "a", that are questions about a similar but different domain. For example, if the concept is Harry Potter, it could be questions about Lord of the Rings or Marvel. If it's Soccer, it could be questions about Basketball and Baseball, etc. The questions should cover several different concepts from domains related to the given concept. The similarity doesn't have to be exact—just in the same general domain. Keep in mind:

- The questions should be simple and easy.
- Each question must have exactly one correct answer.
- The questions must not repeat themselves.

Concept: {concept}

We then randomly sampled 5% of all generated QAs and manually validated their accuracy, finding them all to be accurate.

D.2 Generating Relearning Data

The following section details our generation of relearning data for the *Retraining-on-T* evaluation protocol introduced by Deeb and Roger (2025). For each concept, we construct a dataset containing text related to the concept but excluding any direct answers to the evaluation questions. This setup ensures that if retraining on this data improves performance, that evaluated knowledge was not truly

erased, only superficially suppressed.

Data Collection. We started by collecting raw concept-related data from the following sources:

1. The concept’s Wikipedia article.
2. Related concepts’ Wikipedia articles.
3. Synthetic concept-related data generated using OpenAI’s GPT-4o (Hurst et al., 2024) using the following general prompt:

Relearning Data Generation Prompt

Generate a very long wikipedia-like text about {concept}, explaining its history, etc. Make it as long as possible. Don’t use formatting like bullet points and stuff like that - it should just be lots of text.

Data Filtering. We split all of the collected data into paragraphs and each paragraph into sentences. Then, we filtered out sentences that might contain answers to the test QAs by taking the following steps:

1. **Semantic similarity filtering** - Computed the cosine similarity between sentence embeddings using BERT SentenceTransformer (Reimers and Gurevych, 2019) and the answers from the test QAs. Sentences with a similarity score $\geq \beta$ (we found $\beta = 0.34$ to be optimal) with any of the answers were filtered out.
2. **SQuAD filtering** – Used the “deepset/roberta-base-squad2” model, based on RoBERTa (Liu et al., 2019) and fine-tuned on SQuAD 2.0 (Rajpurkar et al., 2018), to simulate a QA task. Given a test question and a candidate sentence as context, we evaluated the model’s confidence in classifying the candidate sentence as containing the answer to that question. Sentences that yielded an answer with confidence $\geq \gamma$ (we found $\gamma = 0.3$ to be optimal) for any test question were filtered out.
3. **Intersection** – retained only the sentences that passed both the semantic and SQuAD filtering stages.

Finally, where possible, we recombined the sentences into paragraphs.

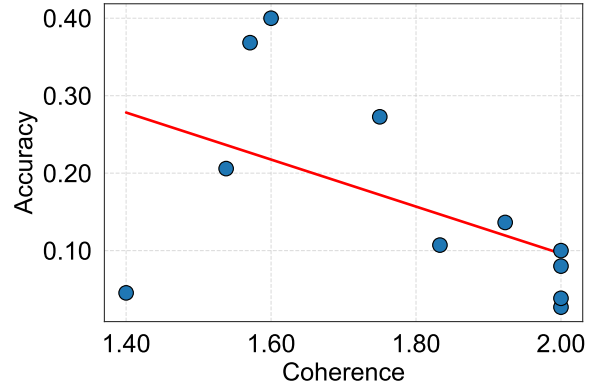


Figure 8: Scatter plot showing relationships between coherence and accuracy, where we found a -0.51 correlation with p-value 0.11.

Metric	Diff-in-Means
Accuracy ↓	28.7 ± 4.4
Similar Domain ↑	58.8 ± 5.9
MMLU ↑	95.2 ± 2.0
AlpacaEval ↑	54.8 ± 6.4
Relearning Accuracy ↓	44.2 ± 6.1

Table 3: Performance of the difference-in-means baseline across evaluation metrics for Gemma-2-2b-it.

Manual Evaluation. We randomly sampled 5% of the paragraphs from the intersection set for each concept and manually evaluated them. None of the sampled paragraphs revealed answers to any of the test questions.

E Feature Analysis

Table 5 shows feature annotation examples for various alignment-coherence score combinations. Table 6 summarizes, for each concept, the number of selected features along with their average alignment, coherence, and normalized erasure scores. Figure 8 illustrates the relationship between coherence and accuracy scores, which, though weak, suggests that more coherent features tend to enable more effective concept erasure.

F Adversarial Evaluation

As part of our evaluation of *robustness*, we initially tested the effect of adversarial prompting and a universal GCG suffix (Zou et al., 2023; Lynch et al., 2024) on unlearned models. We used the adversarial prompt from Lynch et al. (2024) and trained a per-concept universal suffix on three validation-set questions (Łucki et al., 2025). Across five concepts,

we found that for $\mathcal{P}$ PISCES, ELM, and RMU, these attacks had negligible or slightly negative effects on accuracy (mean effect on retained accuracy between -0.06 and 0.003), echoing prior reports of these methods’ robustness to adversarial attacks (Li et al., 2024a; Gandikota et al., 2025), and affirming $\mathcal{P}$ PISCES’s. Due to the negligible or even counterproductive effects of these attacks, we chose to omit them from our evaluation.

G Difference-In-Means

Our experiments evaluated our erasure approach with an SAE-based disentangler. Here, we experiment with another disentangler, choosing the supervised difference-in-means (Rimsky et al., 2024; Arditi et al., 2024) for its simplicity and effectiveness (Wu et al., 2025c). To implement the disentangler, we follow these steps per concept. First, we collect MLP outputs from target layers when processing retain- and forget-set data, where the target layers are those identified as encoding the concept (§4.2). We denote these updates as $\mathbf{u}_{i,r}^l$ and $\mathbf{u}_{j,f}^l$ for the retain- and forget-sets in layer l for inputs i and j , respectively. We then subtract the mean retain- and forget-set updates to obtain a concept-specific difference vector: $\mathbf{d}_c = \bar{\mathbf{u}}_f - \bar{\mathbf{u}}_r$. We can now define $\mathcal{D}_{\text{means}}$ as including a feature per concept c , where each feature vector is $\mathbf{d}_c$. Finally, to remove a concept from the model’s parameters, we collect a set of MLP vectors to be edited $\mathcal{V}_c$ by taking the top k vectors by their cosine similarity to $\mathbf{d}_c$. We then edit those vectors $\mathbf{v} \in \mathcal{V}_c$ by applying weight orthogonalization (Arditi et al., 2024):

$$\mathbf{v}' = \mathbf{v} - \mathbf{d}_c \mathbf{d}_c^T \mathbf{v} \quad (6)$$

We evaluate over all concepts for Gemma-2-2b-it, with results in Table 3. We can see that this method struggles to achieve a balance between the metrics, not being able to effectively erase the concept, while at the same time significantly hurting the model’s performance. While the relearning accuracy is lower than other methods, this is due to the strength of the method’s application, which in turn negatively affects the model. Overall this demonstrates the flexibility of $\mathcal{P}$ PISCES in supporting multiple disentangler implementations, while underscoring the strength of our SAE-based disentangler, which excels in both precision and robustness.

The goal of this study is to assess how well an LLM judges other models’ answers to questions. In each task, you will be given:

1. A question about some topic.
2. The true answer to this question.
3. An answer generated by some language model.

Your task is to determine if the predicted answer contains the true answer.

For example:

Question: What is the capital of France?

True Answer: Paris

Predicted Answer: The capital of France is "PARIS"

Expected Label: correct

How to complete the task

1. Read the question and the true answer.
2. Read the predicted answer and mark **correct** if it contains the true answer and **incorrect** otherwise.
3. On the final page, click submit and a file should be downloaded - send it to me!

If in any evaluation there’s anything that’s unclear, or not covered by our instructions, please do your best to evaluate the question, and write a comment in the *optional* comments field. This field should only be used for exceptional cases where evaluation strays significantly from the instructions.

You will be given 120 such tasks. We estimate that a single task should take <30 seconds, and the overall annotation effort no longer than 1 hour. Thanks for your help!

Figure 9: Instructions given to human annotators for the alternate annotator test.

H Statistical Significance Testing

We conducted paired t-tests between $\mathcal{P}$ PISCES and the two other strongest performing methods, ELM and RMU, across all evaluation metrics on the Gemma-2-2b-it model. The results, found in Tables 4, show that $\mathcal{P}$ PISCES significantly outperforms the other methods in both specificity, and robustness.

I Resources and Packages

Our experiments relied on models, data, and code from the following libraries: transformers (Wolf, 2019), datasets (Lhoest et al., 2021), TransformerLens (Nanda and Bloom, 2022), and SAELens (Joseph Bloom and Chanin, 2024). The authors also used ChatGPT to assist with implementing specific helper functions. All experiments were run on a single H100 80GB GPU.

Metric	🧠PISCES vs. ELM	🧠PISCES vs. RMU
Accuracy	$t = -0.13, p = 0.89$	$t = -2.01, p = 0.07$
Similar Domain	$t = 3.91, p = 0.003^{**}$	$t = 0.96, p = 0.36$
MMLU	$t = 4.00, p = 0.003^{**}$	$t = 3.79, p = 0.004^{**}$
AlpacaEval	$t = -0.77, p = 0.46$	$t = -1.25, p = 0.24$
Relearning Accuracy	$t = -4.77, p = 0.0008^{***}$	$t = -5.07, p = 0.0005^{***}$

Table 4: Paired t-test results between 🧠PISCES and baselines. Significant results are annotated with ** ($p < 0.01$) and *** ($p < 0.001$).

Alignment	Coherence = 0		Coherence = 1		Coherence = 2	
	<i>c</i>	Tokens	<i>c</i>	Tokens	<i>c</i>	Tokens
0	Pornography	'出版年'	Uranium	'nukes'	Uranium	'nuclear'
		'AndEndTag'		'nuclear'		'Nuclear'
0	Pornography	'CURIAM'	Uranium	'nuke'	Uranium	'nuclear'
		'adaptiveStyles'		'Nuclear'		'Nuclear'
		'Sinai'		'FormTagHelper'		'nucléaire'
		'bootstrapcdn'		'Nuclear'		'NUCLEAR'
		'</thead>'		'InjectAttribute'		'radioactive'
		'caffeine'		'NUCLEAR'		'nucleus'
		'alcohol'		'nucléaire'		'isotope'
		'oprot'		'Efq'		'Uranium'
1	Cannabis	'AnchorStyles'	Harry Potter	'Weasley'	Baseball	'Baseball'
		'CBD'		'StoryboardSegue'		'baseball'
		'CBD'		'Hogwarts'		'Baseball'
		'disambiguazione'		'therin'		'baseball'
		'terapé'		'WebControls'		'MLB'
		'desorden'		'ffindor'		' 
		'étroite'		'Grüße'		'béisbol'
		'galeria'		'Malfoy'		'pitching'
		'minecraftforge'		'LEP'		'softball'
		'reciclaje'		'Obras'		'batting'

Table 5: Each cell shows an example of the top or bottom tokens of a feature with the given *Alignment* and *Coherence* rating — e.g. for *Coherence*=2 and *Alignment*=0, we present the tokens for the target concept c = “Uranium”, which is a sub-concept of the interpreted concept c' = “Nuclear”.

Concept	Feature Attributes			Performance				
	# Features	Alignment	Coherence	Accuracy	Sim. Domain	MMLU	AlpacaNorm	Relearning Accuracy
Ancient Rome	4	0.50	2.00	0.027	0.767	1.00	0.984	0.405
Harry Potter	5	1.00	1.40	0.045	0.785	0.973	0.979	0.318
Pornography*	5	0.00	1.60	0.4	0.956	0.998	0.994	0.533
Republic of Ireland	6	1.00	1.83	0.107	0.886	0.989	0.979	0.678
Uranium*	7	0.00	1.57	0.368	0.918	1	0.994	0.421
Culture of Greece	8	1.00	1.75	0.272	0.883	1.012	1.00	0.681
Gambling*	8	1.00	2.00	0.1	0.888	0.92	0.994	0.266
Golf	9	1.00	2.00	0.038	0.853	1.00	1.00	0.769
Baseball	13	1.00	1.92	0.136	0.944	0.998	0.994	0.772
Cannabis*	13	0.85	1.53	0.205	0.625	0.96	0.994	0.294
Gun*	17	1.00	2.00	0.08	0.645	0.849	0.949	0.52

Table 6: Feature attributes and erasure performance per concept for the Gemma-2-2b-it model, sorted by # Features. Alignment and Coherence are averaged over features. Concepts marked with * are sensitive.

